

The background of the book cover is a photograph of a set of concrete stairs. Each step is painted with a different vibrant color, including shades of orange, yellow, green, blue, and purple. A dark metal railing runs along the right side of the stairs. The background is slightly out of focus, showing hints of buildings and foliage.

essentials of
**SOCIAL
STATISTICS**

4e

for a
**Diverse
Society**

Anna Leon-Guerrero
Chava Frankfort-Nachmias
Georgiann Davis



ESSENTIALS OF SOCIAL STATISTICS FOR A DIVERSE SOCIETY

Fourth Edition

Sara Miller McCune founded SAGE Publishing in 1965 to support the dissemination of usable knowledge and educate a global community. SAGE publishes more than 1000 journals and over 800 new books each year, spanning a wide range of subject areas. Our growing selection of library products includes archives, data, case studies and video. SAGE remains majority owned by our founder and after her lifetime will become owned by a charitable trust that secures the company's continued independence.

Los Angeles | London | New Delhi | Singapore | Washington DC | Melbourne

ESSENTIALS OF SOCIAL STATISTICS FOR A DIVERSE SOCIETY

Fourth Edition

Anna Leon-Guerrero

Pacific Lutheran University

Chava Frankfort-Nachmias

University of Wisconsin–Milwaukee

Georgiann Davis

University of Nevada, Las Vegas



Los Angeles | London | New Delhi
Singapore | Washington DC | Melbourne



FOR INFORMATION:

SAGE Publications, Inc.
2455 Teller Road
Thousand Oaks, California 91320
E-mail: order@sagepub.com

SAGE Publications Ltd.
1 Oliver's Yard
55 City Road
London EC1Y 1SP
United Kingdom

SAGE Publications India Pvt. Ltd.
B 1/1 Mohan Cooperative Industrial Area
Mathura Road, New Delhi 110 044
India

SAGE Publications Asia-Pacific Pte. Ltd.
18 Cross Street #10-10/11/12
China Square Central
Singapore 048423

Acquisitions Editor: Helen Salmon
Editorial Assistant: Natalie Elliott
Production Editor: Andrew Olson
Typesetter: Hurix Digital
Proofreader: Larry Baker
Indexer: Integra
Cover Designer: Candice Harman
Marketing Manager: Shari Countryman

Copyright © 2021 by SAGE Publications, Inc.

All rights reserved. No part of this book may be reproduced or utilized in any form or by any means, electronic or mechanical, including photocopying, recording, or by any information storage and retrieval system, without permission in writing from the publisher.

All trademarks depicted within this book, including trademarks appearing as part of a screenshot, figure, or other image are included solely for the purpose of illustration and are the property of their respective holders. The use of the trademarks in no way indicates any relationship with, or endorsement by, the holders of said trademarks. SPSS is a registered trademark of International Business Machines Corporation.

Printed in the United States of America

ISBN: 978-1-5443-7250-1

This book is printed on acid-free paper.

20 21 22 23 24 10 9 8 7 6 5 4 3 2 1

BRIEF CONTENTS

Preface	xv
Acknowledgments	xix
About the Authors	xxi
Statistical Equations	xxii
CHAPTER 1 • The What and the Why of Statistics	1
CHAPTER 2 • The Organization and Graphic Presentation of Data	27
CHAPTER 3 • Measures of Central Tendency and Variability	69
CHAPTER 4 • The Normal Distribution	123
CHAPTER 5 • Sampling and Sampling Distributions	149
CHAPTER 6 • Estimation	175
CHAPTER 7 • Testing Hypotheses	201
CHAPTER 8 • The Chi-Square Test and Measures of Association	237
CHAPTER 9 • Analysis of Variance	279
CHAPTER 10 • Regression and Correlation	303
Appendix A. Table of Random Numbers	349
Appendix B. The Standard Normal Table	353
Appendix C. Distribution of t	359
Appendix D. Distribution of Chi-Square	361
Appendix E. Distribution of F	363
Appendix F. Basic Math Review (on the website*)	
Learning Check Solutions	367

Answers to Odd-Numbered Exercises	381
Glossary	409
Notes	415
Index	421

*Website to be found at: <https://edge.sagepub.com/ssdsess4e>

DETAILED CONTENTS

Preface	xv
Acknowledgments	xix
About the Authors	xxi
Statistical Equations	xxii

CHAPTER 1 • The What and the Why of Statistics 1

The Research Process	2
Asking Research Questions	3
The Role of Theory	4
Formulating the Hypotheses	5
Independent and Dependent Variables: Causality	8
Independent and Dependent Variables: Guidelines	9
Collecting Data	10
Levels of Measurement	10
Nominal Level of Measurement	11
Ordinal Level of Measurement	12
Interval-Ratio Level of Measurement	13
Cumulative Property of Levels of Measurement	13
Levels of Measurement of Dichotomous Variables	14
Discrete and Continuous Variables	16
A Closer Look 1.1: A Cautionary Note: Measurement Error	17
Analyzing Data and Evaluating the Hypotheses	17
Descriptive and Inferential Statistics	18
Evaluating the Hypotheses	19
Examining a Diverse Society	19
Data at Work	21

CHAPTER 2 • The Organization and Graphic Presentation of Data 27

Frequency Distributions	27
Proportions and Percentages	29
Percentage Distributions	30
The Construction of Frequency Distributions	31
Frequency Distributions for Nominal Variables	34
Frequency Distributions for Ordinal Variables	34
Frequency Distributions for Interval-Ratio Variables	35

Cumulative Distributions	38
Rates	40
Bivariate Tables	41
How to Construct a Bivariate Table	41
How to Compute Percentages in a Bivariate Table	44
Calculating Percentages Within Each Category of the Independent Variable	44
Comparing the Percentages Across Different Categories of the Independent Variable	45
Graphic Presentation of Data	46
The Pie Chart	47
The Bar Graph	49
The Histogram	50
The Line Graph	52
The Time-Series Chart	53
Statistics in Practice: Foreign-Born Population 65 Years and Over	54
► A Closer Look 2.1: A Cautionary Note: Distortions in Graphs	56
► Data at Work: Spencer Westby: Senior Editorial Analyst	57

CHAPTER 3 • Measures of Central Tendency and Variability

Measures of Central Tendency	69
The Mode	70
The Median	72
Finding the Median in Sorted Data	72
Finding the Median in Frequency Distributions	76
Locating Percentiles in a Frequency Distribution	77
The Mean	78
► A Closer Look 3.1: Finding the Mean in a Frequency Distribution	81
Understanding Some Important Properties of the Arithmetic Mean	82
Reading the Research Literature: The Case of Reporting Income	85
Statistics in Practice: The Shape of the Distribution	86
The Symmetrical Distribution	86
The Positively Skewed Distribution	87
The Negatively Skewed Distribution	88
Guidelines for Identifying the Shape of a Distribution	90
Considerations for Choosing a Measure of Central Tendency	90
Level of Measurement	90
Skewed Distribution	90

► A Closer Look 3.2: A Cautionary Note: Representing Income	91
<i>Symmetrical Distribution</i>	92
Measures of Variability	92
The Importance of Measuring Variability	93
The Range	94
The Interquartile Range	95
The Box Plot	98
The Variance and the Standard Deviation	101
<i>Calculating the Deviation From the Mean</i>	102
<i>Calculating the Variance and the Standard Deviation</i>	104
Considerations for Choosing a Measure of Variation	106
► A Closer Look 3.3: More on Interpreting the Standard Deviation	108
Reading the Research Literature: Community College Mentoring	110
► Data at Work: Sruthi Chandrasekaran: Senior Research Associate	111

CHAPTER 4 • The Normal Distribution **123**

Properties of the Normal Distribution	123
Empirical Distributions	
Approximating the Normal Distribution	124
Areas Under the Normal Curve	124
Interpreting the Standard Deviation	125
An Application of the Normal Curve	126
Transforming a Raw Score Into a Z Score	127
The Standard Normal Distribution	127
The Standard Normal Table	128
1. Finding the Area Between the Mean and a Positive or Negative Z Score	130
2. Finding the Area Above a Positive Z Score or Below a Negative Z Score	131
3. Transforming Proportions and Percentages Into Z Scores	133
<i>Finding a Z Score That Bounds an Area Above It</i>	133
<i>Finding a Z Score That Bounds an Area Below It</i>	134
4. Working With Percentiles in a Normal Distribution	135
<i>Finding the Percentile Rank of a Score Higher Than the Mean</i>	135
<i>Finding the Percentile Rank of a Score Lower Than the Mean</i>	136
<i>Finding the Raw Score Associated</i>	
<i>With a Percentile Higher Than 50</i>	137
<i>Finding the Raw Score Associated</i>	
<i>With a Percentile Lower Than 50</i>	138

Reading the Research Literature: Child Health and Academic Achievement	140
▶ A Closer Look 4.1: Percentages, Proportions, and Probabilities	140
▶ Data at Work: Claire Wulf Winiarek: Director of Collaborative Policy Engagement	142

CHAPTER 5 • Sampling and Sampling Distributions 149

Aims of Sampling	149
Basic Probability Principles	151
Probability Sampling	153
The Simple Random Sample	154
The Concept of the Sampling Distribution	154
The Population	155
The Sample	156
The Dilemma	157
The Sampling Distribution	157
The Sampling Distribution of the Mean	158
An Illustration	158
Review	160
The Mean of the Sampling Distribution	161
The Standard Error of the Mean	162
The Central Limit Theorem	162
The Size of the Sample	165
The Significance of the Sampling Distribution and the Central Limit Theorem	165
Statistics in Practice: The 2016 U.S. Presidential Election	167
▶ Data at Work: Emily Treichler: Postdoctoral Fellow	168

CHAPTER 6 • Estimation 175

Point and Interval Estimation	176
Confidence Intervals for Means	177
▶ A Closer Look 6.1: Estimation as a Type Inference	178
Determining the Confidence Interval	180
Calculating the Standard Error of the Mean	180
Deciding on the Level of Confidence and Finding the Corresponding Z Value	180
Calculating the Confidence Interval	180
Interpreting the Results	181
Reducing Risk	182
Estimating Sigma	184
Calculating the Estimated Standard Error of the Mean	184
Deciding on the Level of Confidence and Finding the Corresponding Z Value	184

Calculating the Confidence Interval	184
Interpreting the Results	185
Sample Size and Confidence Intervals	185
Statistics in Practice: Hispanic Migration and Earnings	186
► A Closer Look 6.2: What Affects Confidence Interval Width?	189
Confidence Intervals for Proportions	190
Determining the Confidence Interval	192
Calculating the Estimated Standard Error of the Proportion	192
Deciding on the Desired Level of Confidence and Finding the Corresponding Z Value	193
Calculating the Confidence Interval	193
Interpreting the Results	193
Reading the Research Literature: Women Victims of Intimate Violence	194
► Data at Work: Laurel Person Mecca: Research Specialist	196

CHAPTER 7 • Testing Hypotheses 201

Assumptions of Statistical Hypothesis Testing	202
Stating the Research and Null Hypotheses	202
The Research Hypothesis (H_1)	203
The Null Hypothesis (H_0)	204
Probability Values and Alpha	205
► A Closer Look 7.1: More About Significance	208
The Five Steps in Hypothesis Testing: A Summary	209
Errors in Hypothesis Testing	210
The t Statistic and Estimating the Standard Error	211
The t Distribution and Degrees of Freedom	212
Comparing the t and Z Statistics	212
Hypothesis Testing With One Sample and Population Variance Unknown	213
Hypothesis Testing With Two Sample Means	215
The Assumption of Independent Samples	216
Stating the Research and Null Hypotheses	216
The Sampling Distribution of the Difference Between Means	217
Estimating the Standard Error	218
Calculating the Estimated Standard Error	218
The t Statistic	218
Calculating the Degrees of Freedom for a Difference Between Means Test	219
The Five Steps in Hypothesis Testing About Difference Between Means: A Summary	219

► A Closer Look 7.2: Calculating the Estimated Standard Error and the Degrees of Freedom (<i>df</i>) When the Population Variances Are Assumed to Be Unequal	220
Statistics in Practice: Vape Use Among Teens	221
Hypothesis Testing With Two Sample Proportions	223
Reading the Research Literature: Reporting the Results of Hypothesis Testing	227
► Data at Work: Stephanie Wood: Campus Visit Coordinator	229

CHAPTER 8 • The Chi-Square Test and Measures of Association **237**

The Concept of Chi-Square as a Statistical Test	239
The Concept of Statistical Independence	240
The Structure of Hypothesis Testing With Chi-Square	241
The Assumptions	241
Stating the Research and the Null Hypotheses	241
The Concept of Expected Frequencies	241
Calculating the Expected Frequencies	242
Calculating the Obtained Chi-Square	244
The Sampling Distribution of Chi-Square	245
Determining the Degrees of Freedom	246
Making a Final Decision	248
Review	248
Statistics in Practice: Respondent and Mother Education	250
► A Closer Look 8.1: A Cautionary Note: Sample Size and Statistical Significance for Chi-Square	252
Proportional Reduction of Error	253
► A Closer Look 8.2: What Is Strong? What Is Weak? A Guide to Interpretation	254
Lambda: A Measure of Association for Nominal Variables	256
Cramer's <i>V</i> : A Chi-Square-Related Measure of Association for Nominal Variables	259
Gamma and Kendall's Tau- <i>b</i> : Symmetrical Measures of Association for Ordinal Variables	259
Reading the Research Literature:	
India's Internet-Using Population	262
► Data at Work: Patricio Cumsille: Professor	264

CHAPTER 9 • Analysis of Variance **279**

Understanding Analysis of Variance	280
The Structure of Hypothesis Testing With ANOVA	282
The Assumptions	282
Stating the Research and the Null Hypotheses and Setting Alpha	283

The Concepts of Between and Within Total Variance	283
The <i>F</i> Statistic	285
► A Closer Look 9.1: Decomposition of <i>SST</i>	287
Making a Decision	288
The Five Steps in Hypothesis Testing: A Summary	288
Statistics in Practice: The Ethical Consumer	290
► A Closer Look 9.2: Assessing the Relationship Between Variables	291
Reading the Research Literature: College Satisfaction Among Latino Students	291
► Data at Work: Kevin Hemminger: Sales Support Manager/Graduate Program in Research Methods and Statistics	293

CHAPTER 10 • Regression and Correlation 303

The Scatter Diagram	304
Linear Relationships and Prediction Rules	305
Finding the Best-Fitting Line	306
Defining Error	307
► A Closer Look 10.1: Other Regression Techniques	308
The Residual Sum of Squares (Σe^2)	309
The Least Squares Line	309
Computing <i>a</i> and <i>b</i>	309
► A Closer Look 10.2: Understanding the Covariance	312
Interpreting <i>a</i> and <i>b</i>	313
A Negative Relationship: Age and Internet Hours per Week	314
Methods for Assessing the Accuracy of Predictions	317
Calculating Prediction Errors	318
Calculating r^2	322
Testing the Significance of r^2 Using ANOVA	323
Making a Decision	325
Pearson's Correlation Coefficient (<i>r</i>)	326
Characteristics of Pearson's <i>r</i>	326
Statistics in Practice: Multiple Regression and ANOVA	327
► A Closer Look 10.3: Spurious Correlations and Confounding Effects	332
Reading the Research Literature: Academic Intentions and Support	332
► Data at Work: Shinichi Mizokami: Professor	334
Appendix A. Table of Random Numbers	349
Appendix B. The Standard Normal Table	353
Appendix C. Distribution of <i>t</i>	359

Appendix D. Distribution of Chi-Square	361
Appendix E. Distribution of F	363
Appendix F. Basic Math Review (on the website*)	
Learning Check Solutions	367
Answers to Odd-Numbered Exercises	381
Glossary	409
Notes	415
Index	421

*Website to be found at: <https://edge.sagepub.com/ssdsess4e>

PREFACE

You may be reading this introduction on your first day of class. We know you have some questions and concerns about what your course will be like. Math, formulas, and calculations? Yes, those will be part of your learning experience. But there is more.

Throughout our text, we highlight the relevance of statistics in our daily and professional lives. Data are used to predict public opinion, consumer spending, and even a presidential election. How Americans feel about a variety of political and social topics—the Black Lives Matter movement, gun control, immigration, the economy, health care reform, or terrorism—are measured by surveys and polls and reported daily by the news media. Your recent Amazon purchase of hand sanitizer during the COVID-19 pandemic didn't go unnoticed. The study of consumer trends, specifically focusing on young adults, helps determine commercial programming, product advertising and placement, and, ultimately, consumer spending. And as we prepare this text, the world struggles to comprehend political divides around the seriousness of COVID-19, police brutality, migration patterns, and transgender rights, among other things.

Statistics are not just a part of our lives in the form of news bits or information. And it isn't just numbers either. As social scientists, we rely on statistics to help us understand our social world. We use statistical methods and techniques to track demographic trends, to assess social differences, and to better inform social policy. We encourage you to move beyond just being a consumer of statistics and determine how you can use statistics to gain insight into important social issues that affect you and others.

TEACHING AND LEARNING GOALS

Three teaching and learning goals continue to be the guiding principles of our book, as they were in previous editions.

Our first goal is to introduce you to social statistics and demonstrate its value. Although most of you will not use statistics in your own student research, you will be expected to read and interpret statistical information presented by others in professional and scholarly publications, in the workplace, and in the popular media. This book will help you understand the concepts behind the statistics so that you will be able to assess the circumstances in which certain statistics should and should not be used.

A special characteristic of this book is its integration of statistical techniques with substantive issues of particular relevance in the social sciences. Our second goal is to demonstrate that substance and statistical techniques are truly related in social science research. Your learning will not be limited to statistical calculations

and formulas. Rather, you will become proficient in statistical techniques while learning about social differences and inequality through numerous substantive examples and real-world data applications. Because the world we live in is characterized by a growing diversity—where personal and social realities are increasingly shaped by race, class, gender, and other categories of experience—this book teaches you basic statistics while incorporating social science research related to the dynamic interplay of our social worlds.

Our third goal is to enhance your learning by using straightforward prose to explain statistical concepts and by emphasizing intuition, logic, and common sense over rote memorization and derivation of formulas.

DISTINCTIVE AND UPDATED FEATURES OF OUR BOOK

Our learning goals are accomplished through a variety of specific and distinctive features throughout this book.

A Close Link Between the Practice of Statistics, Important Social Issues, and Real-World Examples. This book is distinct for its integration of statistical techniques with pressing social issues of particular concern to society and social science. We emphasize how the conduct of social science is the constant interplay between social concerns and methods of inquiry. In addition, the examples throughout the book—mostly taken from news stories, government reports, public opinion polls, scholarly research, and the National Opinion Research Center’s General Social Survey—are formulated to emphasize to students like you that we live in a world in which statistical arguments are common. Statistical concepts and procedures are illustrated with real data and research, providing a clear sense of how questions about important social issues can be studied with various statistical techniques.

A Focus on Diversity: The United States and International. A strong emphasis on race, class, and gender as central substantive concepts is mindful of a trend in the social sciences toward integrating issues of diversity in the curriculum. This focus on the richness of social differences within our society and our global neighbors is manifested in the application of statistical tools to examine how race, class, gender, and other categories of experience shape our social world and explain social behavior.

Chapter Reorganization and Content. Each revision presents many opportunities to polish and expand the content of our text. In this edition, we have made a number of changes in response to feedback from reviewers and fellow instructors. Measures of central tendency and variability are covered in a single chapter. We also continued to refine our discussion of the interpretation and application of descriptive statistics (variance and standard deviation) and inferential tests (t , Z , F ratio, and regression and correlation). Chapter Exercises do not require the use of computer software.

Reading the Research Literature, Statistics in Practice, A Closer Look, and Data at Work. In your student career and in the workplace, you may be expected

to read and interpret statistical information presented by others in professional and scholarly publications. These statistical analyses are a good deal more complex than most class and textbook presentations. To guide you in reading and interpreting research reports written by social scientists, most of our chapters include a Reading the Research Literature and a Statistics in Practice feature, presenting excerpts of published research reports or specific SPSS calculations using the statistical concepts under discussion. Being statistically literate involves more than just completing a calculation; it also includes learning how to apply and interpret statistical information and being able to say what it means. We include an A Closer Look discussion in each chapter, advising students about the common errors and limitations in quantitative data collection and analysis.

SPSS, Excel, and GSS 2018. IBM® SPSS® Statistics¹ and Microsoft Excel² are used throughout this book, although the use of computers is not required to learn from the text. Real data are used to motivate and make concrete the coverage of statistical topics. As a companion to the fourth edition's SPSS and Excel demonstrations and exercises, we provide three GSS 2018 data sets on the study site at <https://edge.sagepub.com/ssdsess4e>. Two of the GSS data sets (GSS18SSDS-A and GSS18SSDS-B) are formatted for SPSS, while the third data set (GSS18SSDS-E) is ready for use in Excel. These demonstrations and exercises, available on the study site, rely on variables from these modules. There is ample opportunity for instructors to develop their own exercises using these data.

Tools to Promote Effective Study. Each chapter concludes with a list of Main Points and Key Terms discussed in that chapter. Boxed definitions of the Key Terms also appear in the body of the chapter, as do Learning Checks keyed to the most important points. Key Terms are also clearly defined and explained in the Glossary, another special feature in our book. Answers to all the Odd-Numbered Exercises and Learning Checks in the text are included at the end of the book, as well as on the study site at <https://edge.sagepub.com/ssdsess4e>. Complete step-by-step solutions are provided in the Instructor's Manual, available on the study site.

Throughout this text and in ancillary materials, we followed these rounding rules: If the number you are rounding is followed by 5, 6, 7, 8, or 9, round the number up. If the number you are rounding is

followed by 0, 1, 2, 3, or 4, do not change the number. For rounding long decimals, look only at the number in the place you are rounding to and the number that follows it.

¹SPSS is a registered trademark of International Business Machines Corporation.

²Microsoft Excel is a registered trademark of Microsoft Corporation.

On the companion website at: <https://edge.sagepub.com/ssdsess4e>

Instructor site:

- LMS-ready test bank, also available in Microsoft Word
- Editable PowerPoint slides
- Instructors Manual: lecture notes, and answers to all end-of-chapter questions from the book, and the SPSS and Excel problems posted on the student study site. Answers to odd-numbered questions are in the back of the book for students to check their answers.
- Datasets and codebooks
- SPSS and Excel demonstrations and problems to accompany each chapter

Student site:

- eFlashcards of the glossary terms
- Datasets and codebooks
- SPSS and Excel walk-through videos
- SPSS and Excel demonstrations and problems to accompany each chapter
- Appendix F: Basic Math Review

ACKNOWLEDGMENTS

We are grateful to Jeff Lasser and Helen Salmon, publishers for SAGE Publications, for their commitment to our book and for their invaluable assistance through the production process.

We are grateful to Andrew Olson, Kate Russillo, and Gillian Dickens for guiding the book through the production process. We would also like to acknowledge Tiara Beatty, Tara Slagle, Shelly Gupta, and the rest of the SAGE staff for their assistance on this edition.

We are especially indebted to Torisha Khonach, a sociology PhD student at the University of Nevada, Las Vegas, for offering us a fresh pair of eyes throughout the revision process. We deeply appreciate her meticulousness.

Anna Leon-Guerrero would like to thank her Pacific Lutheran University students for inspiring her to be a better teacher. My love and thanks to my husband, Brian Sullivan.

Chava Frankfort-Nachmias would like to thank and acknowledge her friends and colleagues for their unending support; she also would like to thank her students: I am grateful to my students at the University of Wisconsin–Milwaukee, who taught me that even the most complex statistical ideas can be simplified. The ideas presented in this book are the products of many years of classroom testing. I thank my students for their patience and contributions. Finally, I thank my partner, Marlene Stern, for her love and support.

Georgiann Davis would like to thank her students, colleagues, and mentors at the various universities she has been affiliated with over the years for offering her the space to learn and enjoy statistics: I'm especially thankful for my students whose excitement, and sometimes anxiety, continues to nurture my love for quantitative data analysis. A special thank you to my partner, who sometimes shares my love for statistics despite being a hardcore ethnographer.

Anna Leon-Guerrero
Pacific Lutheran University

Chava Frankfort-Nachmias
University of Wisconsin–Milwaukee

Georgiann Davis
University of Nevada, Las Vegas

ABOUT THE AUTHORS

Anna Leon-Guerrero is Professor of Sociology at Pacific Lutheran University in Washington. She received her PhD in sociology from the University of California–Los Angeles. A recipient of the university's Faculty Excellence Award and the K. T. Tang Award for Excellence in Research, she teaches courses in statistics, social theory, and social problems. She is also the author of *Social Problems: Community, Policy, and Social Action*.

Chava Frankfort-Nachmias is an Emeritus Professor of Sociology at the University of Wisconsin–Milwaukee. She is the coauthor of *Research Methods in the Social Sciences* (with David Nachmias), coeditor of *Sappho in the Holy Land* (with Erella Shadmi), and numerous publications on ethnicity and development, urban revitalization, science and gender, and women in Israel. She was the recipient of the University of Wisconsin System teaching improvement grant on integrating race, ethnicity, and gender into the social statistics and research methods curriculum.

Georgiann Davis is Associate Professor of Sociology at the University of Nevada, Las Vegas. An award-winning instructor, researcher, and scholar-activist, she received her PhD in sociology from the University of Illinois at Chicago. She is the author of *Contesting Intersex: The Dubious Diagnosis*.

CHAPTER 5 – THE NORMAL DISTRIBUTION

$$Z = \frac{Y - \bar{Y}}{s}$$

$$Y = \bar{Y} + Z(s)$$

CHAPTER 6 – SAMPLING AND SAMPLING DISTRIBUTIONS

$$K = \frac{\text{Population Size}}{\text{Sample Size}}$$

CHAPTER 7 – ESTIMATION

$$\sigma_{\bar{Y}} = \frac{\sigma}{\sqrt{N}}$$

$$CI = \bar{Y} \pm Z(\sigma_{\bar{Y}})$$

$$s_{\bar{Y}} = \frac{s}{\sqrt{N}}$$

$$\sigma_p = \sqrt{\frac{(\pi)(1-\pi)}{N}}$$

$$s_p = \sqrt{\frac{(p)(1-p)}{N}}$$

$$CI = p \pm Z(s_p)$$

CHAPTER 8 – TESTING HYPOTHESES

$$Z = \frac{\bar{Y} - \mu_Y}{\sigma / \sqrt{N}}$$

$$t = \frac{\bar{Y} - \mu}{s / \sqrt{N}}$$

$$df = N - 1$$

$$\sigma_{\bar{Y}_1 - \bar{Y}_2} = \sqrt{\frac{\sigma_1^2}{N_1} + \frac{\sigma_2^2}{N_2}}$$

$$s_{\bar{Y}_1 - \bar{Y}_2} = \sqrt{\frac{(N_1 - 1)s_1^2 + (N_2 - 1)s_2^2}{(N_1 + N_2) - 2}} \sqrt{\frac{N_1 + N_2}{N_1 N_2}}$$

$$t = \frac{\bar{Y}_1 - \bar{Y}_2}{s_{\bar{Y}_1 - \bar{Y}_2}}$$

$$df = (N_1 + N_2) - 2$$

$$Z = \frac{p_1 - p_2}{S_{p_1 - p_2}}$$

$$s_{p_1 - p_2} = \sqrt{\frac{p_1(1 - p_1)}{N_1} + \frac{p_2(1 - p_2)}{N_2}}$$

CHAPTER 8 – THE CHI-SQUARE TEST AND MEASURES OF ASSOCIATION

$$f_e = \frac{(\text{Column Marginal})(\text{Row Marginal})}{N}$$

$$\chi^2 = \sum \frac{(f_o - f_e)^2}{f_e}$$

$$df = (r - 1)(c - 1)$$

$$\chi^2 = \sum \frac{(|f_o - f_e| - 0.5)^2}{f_e}$$

$$V = \sqrt{\frac{\chi^2}{N(m)}}$$

CHAPTER 9 – ANALYSIS OF VARIANCE

$$SSB = \sum n_k (\bar{Y}_k - \bar{Y})^2$$

$$SSW = \sum (Y_i - \bar{Y}_k)^2$$

$$SSW = \sum Y_i^2 - \sum \frac{(\sum Y_k)^2}{n_k}$$

$$SST = \sum (Y_i - \bar{Y})^2 = SSB + SSW$$

$$df_b = k - 1$$

$$df_w = N - k$$

$$\text{Mean square between} = SSB / df_b$$

$$\text{Mean square within} = SSW / df_w$$

$$F = \frac{\text{Mean square between}}{\text{Mean square within}} = \frac{SSB / df_b}{SSW / df_w}$$

$$\eta^2 = \frac{SSB}{SST}$$

THE WHAT AND THE WHY OF STATISTICS

Are you taking statistics because it is required in your major—not because you find it interesting? If so, you may be feeling intimidated because you associate statistics with numbers, formulas, and abstract notations that seem inaccessible and complicated. Perhaps you feel intimidated not only because you're uncomfortable with math but also because you suspect that numbers and math don't leave room for human judgment or have any relevance to your own personal experience. In fact, you may even question the relevance of statistics to understanding people, social behavior, or society.

In this book, we will show you that statistics can be a lot more interesting and easier to understand than you may have been led to believe. In fact, as we draw on your previous knowledge and experience and relate statistics to interesting and important social issues, you'll begin to see that statistics is not just a course you have to take but a useful tool as well.

There are two reasons why learning statistics may be of value to you. First, you are constantly exposed to statistics every day of your life. Marketing surveys, voting polls, and social research findings appear daily in the news media. By learning statistics, you will become a sharper consumer of statistical material. Second, as a major in the social sciences, you may be expected to read and interpret statistical information related to your occupation or work. Even if conducting research is not a part of your work, you may still be expected to understand and learn from other people's research or to be able to write reports based on statistical analyses.

Just what is statistics, anyway? You may associate the word with numbers that indicate COVID-19 hospitalization rates, support for the Black Lives Matter movement, and so on. But the word **statistics** also refers to a set of procedures used by

Describe the five stages of the research process.

Define independent and dependent variables.

Distinguish between the three levels of measurement.

Apply descriptive and inferential statistical procedures.

A set of procedures used by social scientists to organize, summarize, and communicate numerical information.

Data: Information represented by numbers, which can be the subject of statistical analysis.

Research process: A set of activities in which social scientists engage to answer questions, examine ideas, or test theories.

social scientists to organize, summarize, and communicate numerical information. Only information represented by numbers can be the subject of statistical analysis. Such information is called **data**; researchers use statistical procedures to analyze data to answer research questions and test theories. It is the latter usage—answering research questions and testing theories—that this textbook explores.

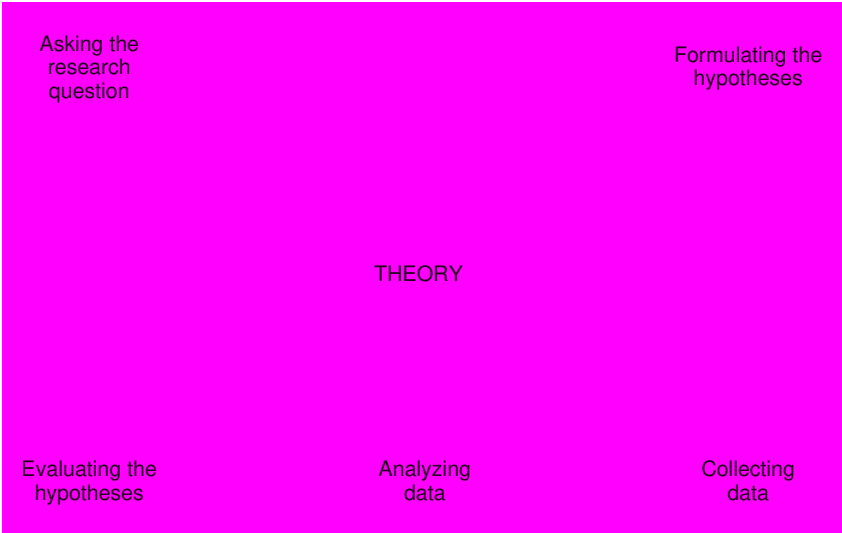
THE RESEARCH PROCESS

To give you a better idea of the role of statistics in social research, let’s start by looking at the **research process**. We can think of the research process as a set of activities in which social scientists engage so that they can answer questions, examine ideas, or test theories.

As illustrated in Figure 1.1, the research process consists of five stages:

1. Asking the research question
2. Formulating the hypotheses
3. Collecting data
4. Analyzing data
5. Evaluating the hypotheses

Figure 1.1 The Research Process



Each stage affects the theory and is affected by it as well. Statistics is most closely tied to the data analysis stage of the research process. As we will see in later chapters, statistical analysis of the data helps researchers test the validity and accuracy of their hypotheses.

ASKING RESEARCH QUESTIONS

The starting point for most research is asking a research question. Consider the following research questions taken from several social science journals:

How does the expansion of police presence in poor urban communities affect educational outcomes?

What does it mean to be a wounded warrior and how does the term impact the way wounded veterans think about themselves?

How do Lebanese women use their informal social networks to engage in political activism for women's rights?

What factors affect the economic mobility of female workers?

These are all questions that can be answered by conducting **empirical research**—research based on information that can be verified by using our direct experience. To answer research questions, we cannot rely on reasoning, speculation, moral judgment, or subjective preference. For example, the questions “Is racial equality good for society?” and “Is an urban lifestyle better than a rural lifestyle?” cannot be answered empirically because the terms good and better are concerned with values, beliefs, or subjective preference and, therefore, cannot be independently verified. One way to study these questions is by defining good and better in terms that can be verified empirically. For example, we can define good in terms of economic growth and better in terms of psychological well-being. These questions could then be answered by conducting empirical research.

You may wonder how to come up with a research question. The first step is to pick a question that interests you. If you are not sure, look around! Ideas for research problems are all around you, from media sources to personal experience or your own intuition. Talk to other people, write down your own observations and ideas, or learn what other social scientists have written about.

Take, for instance, the relationship between gender and work. As a college student about to enter the labor force, you may wonder about the similarities and differences between women's and men's work experiences and about job opportunities when you graduate. Here are some facts and observations based on research reports: In 2018, women who were employed full-time earned about \$794 (in current dollars) per week on average; men who were employed full-time earned \$993 (in current dollars) per week

● **Empirical research:**
A research based on evidence that can be verified by using our direct experience.

on average.¹ Women's and men's work are also very different. Women continue to be the minority in many of the higher-ranking and higher-salaried positions in professional and managerial occupations. For example, in 2017, women made up 18.4% of software developers and 28% of chief executives. In comparison, among all those employed as secretaries and administrative assistants, 96% were women. Among all receptionists and information clerks in 2017, 93% were women.² These observations may prompt us to ask research questions such as the following: How much change has there been in women's work over time? Are women paid, on average, less than men for the same type of work?



LEARNING CHECK 1.1

Identify one or two social science questions amenable to empirical research. You can almost bet that you will be required to do a research project sometime in your college career.

THE ROLE OF THEORY

You may have noticed that each preceding research question was expressed in terms of a relationship. This relationship may be between two or more attributes of individuals or groups, such as gender and income or gender segregation in the workplace and income disparity. The relationship between attributes or characteristics of individuals and groups lies at the heart of social scientific inquiry.

Most of us use the term theory quite casually to explain events and experiences in our daily life. You may have a theory about why your roommate has been so nice to you lately or why you didn't do so well on your last exam. In a somewhat similar manner, social scientists attempt to explain the nature of social reality. Whereas our theories about events in our lives are commonsense explanations based on educated guesses and personal experience, to the social scientist, a theory is a more precise explanation that is frequently tested by conducting research.

A **theory** is a set of assumptions and propositions used by social scientists to explain, predict, and understand the phenomena they study.³ The theory attempts to establish a link between what we observe (the data) and our conceptual understanding of why certain phenomena are related to each other in a particular way.

For instance, suppose we wanted to understand the reasons for the income disparity between men and women; we may wonder whether the types of jobs men and women have and the organizations in which they work

Theory: A set of assumptions and propositions used to explain, predict, and understand social phenomena.

have something to do with their wages. One explanation for gender wage inequality is gender segregation in the workplace—the fact that American men and women are concentrated in different kinds of jobs and occupations. What is the significance of gender segregation in the workplace? In our society, people’s occupations and jobs are closely associated with their level of prestige, authority, and income. The jobs in which women and men are segregated are not only different but also unequal. Although the proportion of women in the labor force has markedly increased, women are still concentrated in occupations with low pay, low prestige, and few opportunities for promotion. Thus, gender segregation in the workplace is associated with unequal earnings, authority, and status. In particular, women’s segregation into different jobs and occupations from those of men is the most immediate cause of the pay gap. Women receive lower pay than men do even when they have the same level of education, skill, and experience as men in comparable occupations.

FORMULATING THE HYPOTHESES

So far, we have come up with several research questions about the income disparity between men and women in the workplace. We have also discussed a possible explanation—a theory—that helps us make sense of gender inequality in wages. Is that enough? Where do we go from here?

Our next step is to test some of the ideas suggested by the gender segregation theory. But this theory, even if it sounds reasonable and logical to us, is too general and does not contain enough specific information to be tested. Instead, theories suggest specific concrete predictions or **hypotheses** about the way that observable attributes of people or groups are interrelated in real life. Hypotheses are tentative because they can be verified only after they have been tested empirically.⁴ For example, one hypothesis we can derive from the gender segregation theory is that wages in occupations in which the majority of workers are female are lower than the wages in occupations in which the majority of workers are male.

Not all hypotheses are derived directly from theories. We can generate hypotheses in many ways—from theories, directly from observations, or from intuition. Probably, the greatest source of hypotheses is the professional or scholarly literature. A critical review of the scholarly literature will familiarize you with the current state of knowledge and with hypotheses that others have studied.

Let’s restate our hypothesis:

Wages in occupations in which the majority of workers are female are lower than the wages in occupations in which the majority of workers are male.

• **Hypothesis:** A statement predicting the relationship between two or more observable attributes.

Variable: A property of people or objects that takes on two or more values.

Unit of analysis: The object of research, such as individuals, groups, organizations, or social artifacts.

Note that this hypothesis is a statement of a relationship between two characteristics that vary: wages and gender composition of occupations. Such characteristics are called variables. A **variable** is a property of people or objects that takes on two or more values. For example, people can be classified into a number of social class categories, such as upper class, middle class, or working class. Family income is a variable; it can take on values from zero to hundreds of thousands of dollars or more. Similarly, gender composition is a variable. The percentage of females (or males) in an occupation can vary from 0 to 100. Wages is a variable, with values from zero to thousands of dollars or more. See Table 1.1 for examples of some variables and their possible values.

Social scientists must also select a **unit of analysis**; that is, they must select the object of their research. We often focus on individual characteristics or behavior, but we could also examine groups of people such as families, formal organizations like elementary schools or corporations, or social artifacts such as children’s books or advertisements. For example, we may be interested in the relationship between an individual’s educational degree and annual income. In this case, the unit of analysis is the individual. On the other hand, in a study of how corporation profits are associated with employee benefits, corporations are the unit of analysis. If we examine how often women are featured in prescription drug advertisements, the advertisements are the unit of analysis. Figure 1.2 illustrates different units of analysis frequently employed by social scientists.

Table 1.1 Variables and Value Categories	
Variable	Categories
Social class	Lower Working Middle Upper
Gender	Male Female
Education	Less than high school High school Some college College graduate

Figure 1.2 Examples of Units of Analysis

Individual as unit of analysis:

How old are you?
What are your political views?
What is your occupation?

Family as unit of analysis:

How many children are in the family?
Who does the housework?
How many wage earners are there?



Organization as unit of analysis:

How many employees are there?
What is the gender composition?
Do you have a diversity office?

City as unit of analysis:

What was the crime rate last year?
What is the population density?
What type of government runs things?

LEARNING CHECK 1.2

Remember that research question you came up with? Formulate a testable hypothesis based on your research question. Remember that your variables must take on two or more values and you must determine the unit of analysis. What is your unit of analysis?



Independent and Dependent Variables: Causality

Hypotheses are usually stated in terms of a relationship between an independent and a dependent variable. The distinction between an independent and a dependent variable is important in the language of research. Social theories often intend to provide an explanation for social patterns or causal relations between variables. For example, according to the gender segregation theory, gender segregation in the workplace is the primary explanation (although certainly not the only one) of the male-female earning gap. Why should jobs where the majority of workers are women pay less than jobs that employ mostly men? One explanation is that

societies undervalue the work women do, regardless of what those tasks are, because women do them. . . . For example, our culture tends to devalue caring or nurturant work at least partly because women do it. This tendency accounts for childcare workers' low rank in the pay hierarchy.⁵

Dependent variable:

The variable to be explained (the "effect").

Independent variable:

The variable expected to account for (the "cause" of) the dependent variable.

In the language of research, the variable the researcher wants to explain ("the effect") is called the **dependent variable**. The variable that is expected to "cause" or account for the dependent variable is called the **independent variable**. Therefore, in our example, *gender composition of occupations* is the independent variable, and *wages* is the dependent variable.

Cause-and-effect relationships between variables are not easy to infer in the social sciences. To establish that two variables are causally related, your analysis must meet three conditions: (1) The cause has to precede the effect in time, (2) there has to be an empirical relationship between the cause and the effect, and (3) this relationship cannot be explained by other factors.

Let's consider the decades-old debate about controlling crime through the use of prevention versus punishment. Some people argue that special counseling for youths at the first sign of trouble and strict controls on access to firearms would help reduce crime. Others argue that overhauling federal and state sentencing laws to stop early prison releases is the solution. In the early 1990s, Washington and California adopted "three strikes and you're out" legislation, imposing life prison terms on three-time felony offenders. Such laws are also referred to as habitual or persistent offender laws. Twenty-six other states and the federal government adopted similar measures, all advocating a "get tough" policy on crime; the most recent legislation was in 2012 in the state of Massachusetts. In 2012, California voters supported a revision to the original law, imposing a life sentence only when the new felony conviction is serious or violent. Let's suppose that years after the measure was introduced, the crime rate declined in some of these states (in fact, advocates of the measure have identified declining crime rates as evidence of its success). Does the observation that the incidence of crime declined mean that the new measure caused this reduction? Not necessarily! Perhaps the rate of crime had been going down for other reasons, such as improvement in the economy, and the new measure had nothing to do with it. To demonstrate a cause-and-effect relationship,

we would need to show three things: (1) The reduction of crime actually occurred after the enactment of this measure, (2) the enactment of the “three strikes and you’re out” measure was empirically associated with a decrease in crime, and (3) the relationship between the reduction in crime and the “three strikes and you’re out” policy is not due to the influence of another variable (e.g., the improvement of overall economic conditions).

Independent and Dependent Variables: Guidelines

Because it is difficult to infer cause-and-effect relationships in the social sciences, be cautious about using the terms cause and effect when examining relationships between variables. However, using the terms independent variable and dependent variable is still appropriate even when this relationship is not articulated in terms of direct cause and effect. Here are a few guidelines that may help you identify the independent and dependent variables:

1. The dependent variable is always the property that you are trying to explain; it is always the object of the research.
2. The independent variable usually occurs earlier in time than the dependent variable.
3. The independent variable is often seen as influencing, directly or indirectly, the dependent variable.

The purpose of the research should help determine which is the independent variable and which is the dependent variable. In the real world, variables are neither dependent nor independent; they can be switched around depending on the research problem. A variable defined as independent in one research investigation may be a dependent variable in another.⁶ For instance, *educational attainment* may be an independent variable in a study attempting to explain how education influences political attitudes. However, in an investigation of whether a person’s level of education is influenced by the social status of his or her family of origin, *educational attainment* is the dependent variable. Some variables, such as race, age, and ethnicity, because they are primordial characteristics that cannot be explained by social scientists, are never considered dependent variables in a social science analysis.

LEARNING CHECK 1.3

Identify the independent and dependent variables in the following hypotheses:

Older Americans are more likely to support stricter immigration laws than younger Americans.

(Continued)



(Continued)

People who attend church regularly are more likely to oppose abortion than people who do not attend church regularly.

Elderly women are more likely to live alone than elderly men.

Individuals with postgraduate education are likely to have fewer children than those with less education.

What are the independent and dependent variables in your hypothesis?

COLLECTING DATA

Once we have decided on the research question, the hypothesis, and the variables to be included in the study, we proceed to the next stage in the research cycle. This step includes measuring our variables and collecting the data. As researchers, we must decide how to measure the variables of interest to us, how to select the cases for our research, and what kind of data collection techniques we will be using. A wide variety of data collection techniques are available to us, from direct observations to survey research, experiments, or secondary sources. Similarly, we can construct numerous measuring instruments. These instruments can be as simple as a single question included in a questionnaire or as complex as a composite measure constructed through the combination of two or more questionnaire items. The choice of a particular data collection method or instrument to measure our variables depends on the study objective. For instance, suppose we decide to study how one's social class is related to attitudes about women in the labor force. Since attitudes about working women are not directly observable, we need to collect data by asking a group of people questions about their attitudes and opinions. A suitable method of data collection for this project would be a survey that uses a questionnaire or interview guide to elicit verbal reports from respondents. The questionnaire could include numerous questions designed to measure attitudes toward working women, social class, and other variables relevant to the study.

How would we go about collecting data to test the hypothesis relating the gender composition of occupations to wages? We want to gather information on the proportion of men and women in different occupations and the average earnings for these occupations. This kind of information is routinely collected and disseminated by the U.S. Department of Labor, the Bureau of Labor Statistics, and the U.S. Census Bureau. We could use these data to test our hypothesis.

Levels of Measurement

The statistical analysis of data involves many mathematical operations, from simple counting to addition and multiplication. However, not every operation

can be used with every variable. The type of statistical operation we employ depends on how our variables are measured. For example, for the variable *gender*, we can use the number 1 to represent females and the number 2 to represent males. Similarly, 1 can also be used as a numerical code for the category “one child” in the variable *number of children*. Clearly, in the first example, the number is an arbitrary symbol that does not correspond to the property “female,” whereas in the second example, the number 1 has a distinct numerical meaning that does correspond to the property “one child.” The correspondence between the properties we measure and the numbers representing these properties determines the type of statistical operations we can use. The degree of correspondence also leads to different ways of measuring—that is, to distinct levels of measurement. In this section, we will discuss three levels of measurement: (1) nominal, (2) ordinal, and (3) interval-ratio.

Nominal Level of Measurement

With a **nominal level of measurement**, numbers or other symbols are assigned a set of categories for the purpose of naming, labeling, or classifying the observations. *Gender* is an example of a nominal-level variable (Table 1.2). Using the numbers 1 and 2, for instance, we can classify our observations into the categories “females” and “males,” with 1 representing females and 2 representing males. We could use any of a variety of symbols to represent the different categories of a nominal variable; however, when numbers are used to represent the different categories, we do not imply anything about the magnitude or quantitative difference between the categories. Nominal categories cannot be rank-ordered. Because the different categories (e.g., males vs. females) vary in the quality inherent in each but not in quantity, nominal variables are often

● **Nominal level of measurement:** Numbers or other symbols are assigned to a set of categories for the purpose of naming, labeling, or classifying the observations. Nominal categories cannot be rank-ordered.

Table 1.2 Nominal Variables and Value Categories	
Variable	Categories
Gender	Male Female
Religion	Protestant Christian Jewish Muslim
Marital status	Married Single Widowed Other

called qualitative. Other examples of nominal-level variables are political party, religion, and race.

Nominal variables should include categories that are both exhaustive and mutually exclusive. Exhaustiveness means that there should be enough categories composing the variables to classify every observation. For example, the common classification of the variable *marital status* into the categories “married,” “single,” and “widowed” violates the requirement of exhaustiveness. As defined, it does not allow us to classify same-sex couples or heterosexual couples who are not legally married. We can make every variable exhaustive by adding the category “other” to the list of categories. However, this practice is not recommended if it leads to the exclusion of categories that have theoretical significance or a substantial number of observations.

Mutual exclusiveness means that there is only one category suitable for each observation. For example, we need to define religion in such a way that no one would be classified into more than one category. For instance, the categories Protestant and Methodist are not mutually exclusive because Methodists are also considered Protestant and, therefore, could be classified into both categories.



LEARNING CHECK 1.4

Review the definitions of exhaustive and mutually exclusive. Now look at Table 1.2. What other categories could be added to each variable to be exhaustive and mutually exclusive?

Ordinal level of measurement: Numbers are assigned to rank-ordered categories ranging from low to high or high to low.

Ordinal Level of Measurement

Whenever we assign numbers to rank-ordered categories ranging from low to high or high to low, we have an **ordinal level of measurement**. *Social class* is an example of an ordinal variable. We might classify individuals with respect to their social class status as “upper class,” “middle class,” or “working class.” We can say that a person in the category “upper class” has a higher class position than a person in a “middle-class” category (or that a “middle-class” position is higher than a “working-class” position), but we do not know the magnitude of the differences between the categories—that is, we don’t know how much higher “upper class” is compared with the “middle class.”

Many attitudes that we measure in the social sciences are ordinal-level variables. Take, for instance, the following statement used to measure attitudes toward working women: “Women should return to their traditional role in society.” Respondents are asked to identify the number representing their degree of agreement or disagreement with this statement. One form in which a number might be made to correspond with the answers can be seen in

Table 1.3 Ordinal Ranking Scale	
Rank	Value
1	Strongly agree
2	Agree
3	Neither agree nor disagree
4	Disagree
5	Strongly disagree

Table 1.3. Although the differences between these numbers represent higher or lower degrees of agreement with the statement, the distance between any two of those numbers does not have a precise numerical meaning.

Like nominal variables, ordinal variables should include categories that are mutually exhaustive and exclusive.

Interval-Ratio Level of Measurement

If the categories (or values) of a variable can be rank-ordered and if the measurements for all the cases are expressed in the same units and equally spaced, then an **interval-ratio level of measurement** has been achieved. Examples of variables measured at the interval-ratio level are *age*, *income*, and *SAT scores*. With all these variables, we can compare values not only in terms of which is larger or smaller but also in terms of how much larger or smaller one is compared with another. In some discussions of levels of measurement, you will see a distinction made between interval-ratio variables that have a natural zero point (where zero means the absence of the property) and those variables that have zero as an arbitrary point. For example, weight and length have a natural zero point, whereas temperature has an arbitrary zero point. Variables with a natural zero point are also called *ratio variables*. In statistical practice, however, ratio variables are subjected to operations that treat them as interval and ignore their ratio properties. Therefore, we make no distinction between these two types in this text.

● **Interval-ratio level of measurement:** Measurements for all cases are expressed in the same units and equally spaced. Interval-ratio values can be rank-ordered.

Cumulative Property of Levels of Measurement

Variables that can be measured at the interval-ratio level of measurement can also be measured at the ordinal and nominal levels. As a rule, properties that can be measured at a higher level (interval-ratio is the highest) can also be measured at lower levels, but not vice versa. Let's take, for example, *gender composition of occupations*, the independent variable in our research example. Table 1.4 shows the percentage of women in five major occupational groups.

Table 1.4 Gender Composition of Five Major Occupational Groups, 2018	
Occupational Group	Women in Occupation (%)
Management, professional, and related occupations	51.5
Service occupations	57.5
Production, transportation, and materials occupations	23.1
Sales and office occupations	61.1
Natural resources, construction, and maintenance occupations	5.1

Source: U.S. Department of Labor, 2018, Labor Force Statistics from the Current Population Survey 2018, Table 11.

The variable *gender composition* (measured as the percentage of women in the occupational group) is an interval-ratio variable and, therefore, has the properties of nominal, ordinal, and interval-ratio measures. For example, we can say that the management group differs from the natural resources group (a nominal comparison), that service occupations have more women than the other occupational categories (an ordinal comparison), and that service occupations have 34.4 percentage points more women (57.5–23.1) than production occupations (an interval-ratio comparison).

The types of comparisons possible at each level of measurement are summarized in Table 1.5 and Figure 1.3. Note that differences can be established at each of the three levels, but only at the interval-ratio level can we establish the magnitude of the difference.

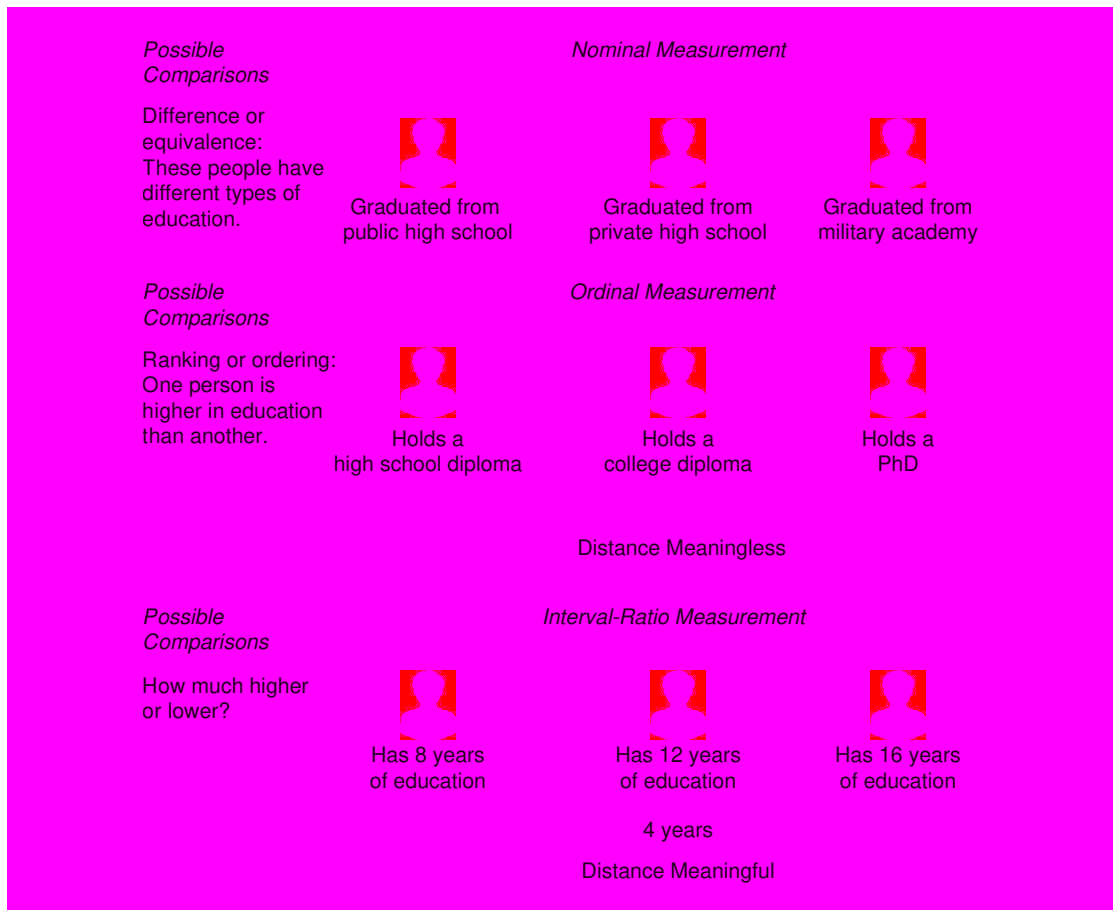
Levels of Measurement of Dichotomous Variables

Dichotomous variable:
A variable that has only two values.

A variable that has only two values is called a **dichotomous variable**. Several key social factors, such as gender, employment status, and marital status, are dichotomies—that is, you are male or female, employed or unemployed, married or not married. Such variables may seem to be measured at the nominal level: You fit in either one category or the other. No category is naturally higher or lower than the other, so they can’t be ordered.

Table 1.5 Levels of Measurement and Possible Comparisons			
Level	Different or Equivalent	Higher or Lower	How Much Higher
Nominal	Yes	No	No
Ordinal	Yes	Yes	No
Interval-ratio	Yes	Yes	Yes

Figure 1.3 Levels of Measurement and Possible Comparisons: Education Measured on Nominal, Ordinal, and Interval-Ratio Levels



However, because there are only two possible values for a dichotomy, we can measure it at the ordinal or the interval-ratio level. For example, we can think of “femaleness” as the ordering principle for gender, so that “female” is higher and “male” is lower. Using “maleness” as the ordering principle, “female” is lower and “male” is higher. In either case, with only two classes, there is no way to get them out of order; therefore, gender could be considered at the ordinal level.

Dichotomous variables can also be considered to be interval-ratio level. Why is this? In measuring interval-ratio data, the size of the interval between the categories is meaningful: The distance between 4 and 7, for example, is the same as the distance between 11 and 14. But with a dichotomy, there is only one interval. Therefore, there is really no other distance to which we can compare it.

Mathematically, this gives the dichotomy more power than other nominal-level variables (as you will notice later in the text).

For this reason, researchers often dichotomize some of their variables, turning a multicategory nominal variable into a dichotomy. For example, you may see race dichotomized into “white” and “nonwhite.” Though we would lose the ability to examine each unique racial category and we may collapse categories that are not similar, it may be the most logical statistical step to take. When you dichotomize a variable, be sure that the two categories capture a distinction that is important to your research question (e.g., a comparison of the number of white vs. nonwhite U.S. senators).



LEARNING CHECK 1.5

Make sure you understand these levels of measurement. As the course progresses, your instructor is likely to ask you what statistical procedure you would use to describe or analyze a set of data. To make the proper choice, you must know the level of measurement of the data.

Discrete and Continuous Variables

The statistical operations we can perform are also determined by whether the variables are continuous or discrete. Discrete variables have a minimum-sized unit of measurement, which cannot be subdivided. The number of children per family is an example of a discrete variable because the minimum unit is one child. A family may have two or three children, but not 2.5 children. The variable *wages* in our research example is a discrete variable because currency has a minimum unit (1 cent), which cannot be subdivided. One can have \$101.21 or \$101.22 but not \$101.21843. Wages cannot differ by less than 1 cent—the minimum-sized unit.

Unlike discrete variables, continuous variables do not have a minimum-sized unit of measurement; their range of values can be subdivided into increasingly smaller fractional values. *Length* is an example of a continuous variable because there is no minimum unit of length. A particular object may be 12 in. long, it may be 12.5 in. long, or it may be 12.532011 in. long. Although we cannot always measure all possible length values with absolute accuracy, it is possible for objects to exist at an infinite number of lengths.⁷ In principle, we can speak of a tenth of an inch, a ten thousandth of an inch, or a ten trillionth of an inch. The variable *gender composition of occupations* is a continuous variable because it is measured in proportions or percentages (e.g., the percentage of women civil engineers), which can be subdivided into smaller and smaller fractions.

This attribute of variables—whether they are continuous or discrete—affects subsequent research operations, particularly measurement procedures,

data analysis, and methods of inference and generalization. However, keep in mind that, in practice, some discrete variables can be treated as if they were continuous, and vice versa.

LEARNING CHECK 1.6

Name three continuous and three discrete variables. Determine whether each of the variables in your hypothesis is continuous or discrete.



A Cautionary Note: Measurement Error

Social scientists attempt to ensure that the research process is as error free as possible, beginning with how we construct our measurements. We pay attention to two characteristics of measurement: (1) reliability and (2) validity.

Reliability means that the measurement yields consistent results each time it is used. For example, asking a sample of individuals, “Do you approve or disapprove of President Donald Trump’s job performance?” is more reliable than asking “What do you think of President Donald Trump’s job performance?” While responses to the second question are meaningful, the answers might be vague and could be subject to different interpretations. Researchers look for the consistency of measurement over time, in relationship with other

related measures, or in measurements or observations made by two or more researchers. Reliability is a prerequisite for validity: We cannot measure a phenomenon if the measure we are using gives us inconsistent results.

Validity refers to the extent to which measures indicate what they are intended to measure. While standardized IQ tests are reliable, it is still debated whether such tests measure intelligence or one’s test-taking ability. A measure may not be valid due to individual error (individuals may want to provide socially desirable responses) or method error (questions may be unclear or poorly written).

Specific techniques and practices for determining and improving measurement reliability and validity are the subject of research methods courses.

ANALYZING DATA AND EVALUATING THE HYPOTHESES

Following the data collection stage, researchers analyze their data and evaluate the hypotheses of the study. The data consist of codes and numbers used to represent their observations. In our example, two scores would represent each occupational group: (1) the percentage of women and (2) the average wage. If we had collected information on 100 occupations, we would end up with 200

scores, 2 per occupational group. However, the typical research project includes more variables; therefore, the amount of data the researcher confronts is considerably larger. We now must find a systematic way to organize these data, analyze them, and use some set of procedures to decide what they mean. These last steps make up the statistical analysis stage, which is the main topic of this textbook. It is also at this point in the research cycle that statistical procedures will help us evaluate our research hypothesis and assess the theory from which the hypothesis was derived.

Descriptive and Inferential Statistics

Statistical procedures can be divided into two major categories: (1) descriptive statistics and (2) inferential statistics. Before we can discuss the difference between these two types of statistics, we need to understand the terms population and sample. A **population** is the total set of individuals, objects, groups, or events in which the researcher is interested. For example, if we were interested in looking at voting behavior in the last presidential election, we would probably define our population as all citizens who voted in the election. If we wanted to understand the employment patterns of Latinas in our state, we would include in our population all Latinas in our state who are in the labor force.

Although we are usually interested in a population, quite often, because of limited time and resources, it is impossible to study the entire population. Imagine interviewing all the citizens of the United States who voted in the last election or even all the Latinas who are in the labor force in our state. Not only would that be very expensive and time-consuming, but we would also probably have a very hard time locating everyone! Fortunately, we can learn a lot about a population if we carefully select a subset from that population. A subset of cases selected from a population is called a **sample**. The process of identifying and selecting this subset is referred to as **sampling**. Researchers usually collect their data from a sample and then generalize their observations to the population. The ultimate goal of sampling is to have a subset that closely resembles the characteristics of the population. Because the sample is intended to represent the population that we are interested in, social scientists take sampling seriously. We'll explore different sampling methods in Chapter 5.

Descriptive statistics includes procedures that help us organize and describe data collected from either a sample or a population. Occasionally, data are collected on an entire population, as in a census. **Inferential statistics**, on the other hand, make predictions or inferences about a population based on observations and analyses of a sample. For instance, the General Social Survey (GSS), from which numerous examples presented in this book are drawn, is conducted every other year by the National Opinion Research Center (NORC) on a representative sample of several thousands of respondents. The survey, which includes several hundred questions (the data collection interview takes approximately 90 minutes), is designed to provide social science researchers with a readily accessible database of socially relevant attitudes, behaviors, and

Population: The total set of individuals, objects, groups, or events in which the researcher is interested.

Sample: A subset of cases selected from a population.

Sampling: The process of identifying and selecting the subset of the population for study.

Descriptive statistics: Procedures that help us organize and describe data collected from either a sample or a population.

Inferential statistics: The logic and procedures concerned with making predictions or inferences about a population from observations and analyses of a sample.

attributes of a cross section of the U.S. adult (18 years of age or older) population. Since 2006, the survey has been administered in English and Spanish. NORC has verified that the composition of the GSS samples closely resembles census data. But because the data are based on a sample rather than on the entire population, the average of the sample does not equal the average of the population as a whole.

Evaluating the Hypotheses

At the completion of these descriptive and inferential procedures, we can move to the next stage of the research process: the assessment and evaluation of our hypotheses and theories in light of the analyzed data. At this next stage, new questions might be raised about unexpected trends in the data and about other variables that may have to be considered in addition to our original variables. For example, we may have found that the relationship between gender composition of occupations and earnings can be observed with respect to some groups of occupations but not others. Similarly, the relationship between these variables may apply for some racial/ethnic groups but not for others.

These findings provide evidence to help us decide how our data relate to the theoretical framework that guided our research. We may decide to revise our theory and hypothesis to take account of these later findings. Recent studies are modifying what we know about gender segregation in the workplace. These studies suggest that race as well as gender shape the occupational structure in the United States and help explain disparities in income. This reformulation of the theory calls for a modified hypothesis and new research, which starts the circular process of research all over again.

Statistics provides an important link between theory and research. As our example on gender segregation demonstrates, the application of statistical techniques is an indispensable part of the research process. The results of statistical analyses help us evaluate our hypotheses and theories, discover unanticipated patterns and trends, and provide the impetus for shaping and reformulating our theories. Nevertheless, the importance of statistics should not diminish the significance of the preceding phases of the research process. Nor does the use of statistics lessen the importance of our own judgment in the entire process. Statistical analysis is a relatively small part of the research process, and even the most rigorous statistical procedures cannot speak for themselves. If our research questions are poorly conceived or our data are flawed due to errors in our design and measurement procedures, our results will be useless.

EXAMINING A DIVERSE SOCIETY

The increasing diversity of American society is relevant to social science. By the middle of this century, if current trends continue unchanged, the United States will no longer be comprised predominantly of European immigrants and

their descendants. Due mostly to renewed immigration and higher birthrates, in time, nearly half the U.S. population will be of African, Asian, Latinx, or Native American ancestry.

Less partial and distorted explanations of social relations tend to result when researchers, research participants, and the research process itself reflect that diversity. A consciousness of social differences shapes the research questions we ask, how we observe and interpret our findings, and the conclusions we draw. Although diversity has been traditionally defined by race, class, and gender, other social characteristics such as sexual identity, physical ability, religion, and age have been identified as important dimensions of diversity. Statistical procedures and quantitative methodologies can be used to describe our diverse society, and we will begin to look at some applications in the next chapter. For now, we will preview some of these statistical procedures.

In Chapter 2, we will learn how to organize information using descriptive statistics, frequency distributions, bivariate tables, and graphic techniques. These statistical tools can also be employed to learn about the characteristics and experiences of groups in our society that have not been as visible as other groups. For example, in a series of special reports published by the U.S. Census Bureau over the past few years, these descriptive statistical techniques have been used to describe the characteristics and experiences of ethnic minorities and those who are foreign born. Using data published by the U.S. Census Bureau, we discuss various graphic devices that can be used to explore the differences and similarities among the many social groups coexisting within the American society. These devices are also used to emphasize the changing age composition of the U.S. population.

In Chapter 3, we describe how to calculate and describe the similarities and commonalities in social experiences (measures of central tendency) and the differences and diversity within social groups (measures of variability). We examine a variety of social demographic variables, including the ethnic composition of the 50 U.S. states.

We will learn about inferential statistics and bivariate analyses in Chapters 4 through 10. First, we review the bases of inferential statistics—the normal distribution, sampling and probability, and estimation—in Chapters 4 to 6. In Chapters 7 to 10, we examine the ways in which class, sex, and ethnicity influence various social behaviors and attitudes. Inferential statistics, such as the t test, chi-square, and the F statistic, help us determine the error involved in using our samples to answer questions about the population from which they are drawn. In addition, we review several methods of bivariate analysis, which are especially suited for examining the association between different social behaviors and attitudes and variables such as race, class, ethnicity, gender, and religion. We use these methods of analysis to show not only how each of these variables operates independently in shaping behavior but also how they interlock to shape our experience as individuals in society.⁸

Whichever model of social research you use—whether you follow a traditional one or integrate your analysis with qualitative data, whether you focus on social differences or any other aspect of social behavior—remember that any application of statistical procedures requires a basic understanding of the statistical concepts and techniques. This introductory text is intended to familiarize you with the range of descriptive and inferential statistics widely applied in the social sciences. Our emphasis on statistical techniques should not diminish the importance of human judgment and your awareness of the person-made quality of statistics. Only with this awareness can statistics become a useful tool for understanding diversity and social life.

At the end of each chapter, the Data at Work feature will introduce you to people who use quantitative data and research methods in their professional lives. They represent a wide range of career fields—education, clinical psychology, international studies, public policy, publishing, politics, and research. Some may have been led to their current positions because of the explicit integration of quantitative

data and research, while others are accidental data analysts—quantitative data became part of their work portfolio. Although “data” or “statistics” are not included in their job titles, these individuals are collecting, disseminating, and/or analyzing data.

We encourage you to review each profile and imagine how you could use quantitative data and methods at work.

DATA AT WORK

MAIN POINTS

Social scientists use statistics to organize, summarize, and communicate information. Only information represented by numbers can be the subject of statistical analysis.

The research process is a set of activities in which social scientists engage to answer questions, examine ideas, or test theories. It consists of the following stages: asking the research question,

formulating the hypotheses, collecting data, analyzing data, and evaluating the hypotheses.

A theory is a set of assumptions and propositions used for explanation, prediction, and understanding of social phenomena. Theories offer specific concrete predictions about the way observable attributes of people or groups would be interrelated in real

life. These predictions, called hypotheses, are tentative answers to research problems.

A variable is a property of people or objects that takes on two or more values. The variable that the researcher wants to explain (the “effect”) is called the dependent variable. The variable that is expected to “cause” or account for the dependent variable is called the independent variable.

Three conditions are required to establish causal relations: (1) The cause has to precede the effect in time, (2) there has to be an empirical relationship between the cause and the effect, and (3) this relationship cannot be explained by other factors.

At the nominal level of measurement, numbers or other symbols are assigned to a set of categories to name, label, or

classify the observations. At the ordinal level of measurement, categories can be rank-ordered from low to high (or vice versa). At the interval-ratio level of measurement, measurements for all cases are expressed in the same unit.

A population is the total set of individuals, objects, groups, or events in which the researcher is interested. A sample is a relatively small subset selected from a population. Sampling is the process of identifying and selecting the subset.

Descriptive statistics includes procedures that help us organize and describe data collected from either a sample or a population. Inferential statistics is concerned with making predictions or inferences about a population from observations and analyses of a sample.

KEY TERMS

data	2	independent variable	8	population	18
dependent variable	8	inferential statistics	18	research process	2
descriptive		interval-ratio level of		sample	18
statistics	18	measurement	13	sampling	18
dichotomous		nominal level of		statistics	1
variable	14	measurement	11	theory	4
empirical research	3	ordinal level of		unit of analysis	6
hypothesis	5	measurement	12	variable	6

DIGITAL RESOURCES

Access key study tools at <https://edge.sagepub.com/ssdsess4e>

eFlashcards of the glossary terms

Datasets and codebooks

SPSS and Excel walk-through videos

SPSS and Excel demonstrations and problems to accompany each chapter

Appendix F: Basic Math Review

CHAPTER EXERCISES

In your own words, explain the relationship of data (collecting and analyzing) to the research process. (Refer to Figure 1.1.)

Construct potential hypotheses or research questions to relate the variables in each of the following examples. Also, write a brief statement explaining why you believe there is a relationship between the variables as specified in your hypotheses.

Political party and support of a U.S.–Mexico border wall

Income and race/ethnicity

The crime rate and the number of police in a city

Life satisfaction and marital status

Age and support for marijuana legalization

Care of elderly parents and ethnicity

Determine the level of measurement for each of the following variables:

The number of people in your statistics class

The percentage of students who are first-generation college students at your school

The name of each academic major offered in your college

The level of support for the Black Lives Matter movement, on a scale from “strong support” to “no support.”

The type of transportation a person takes to school (e.g., bus, walk, car)

The percentage of community members who tested positive for COVID-19

The rating of the overall quality of your campus coffee shop, on a scale from “excellent” to “poor”

For each of the variables in Exercise 3 that you classified as interval-ratio, identify whether it is discrete or continuous.

Why do you think men and women, on average, do not earn the same amount of money? Develop your own theory to explain the difference. Use three independent variables in your theory, with annual income as your dependent variable. Construct hypotheses to link each independent variable with your dependent variable.

For each of the following examples, indicate whether it involves the use of descriptive or inferential statistics. Justify your answer.

The number of unemployed people in the United States

Determining students' opinion about the quality of food at the cafeteria based on a sample of 100 students

The national incidence of breast cancer among Asian women

Conducting a study to determine the rating of the quality of a new smartphone, gathered from 1,000 new buyers

The average GPA of various majors (e.g., sociology, psychology, English) at your university

The change in the number of immigrants coming to the United States from Southeast Asian countries between 2010 and 2015

Adela García-Aracil (2007)⁹ identified how several factors affected the earnings of young European higher-education graduates. Based on data from several EU (European Union) countries, her statistical models included the following variables: annual income (actual dollars), gender (male or female), the number of hours worked per week (actual hours), and years of education (actual years) for each graduate. She also identified each graduate by current job title (senior officials and managers, professionals, technicians, clerks, or service workers).

What is García-Aracil's dependent variable?

Identify two independent variables in her research. Identify the level of measurement for each.

Based on her research, García-Aracil can predict the annual income for other young graduates with similar work experiences and characteristics like the graduates in her sample. Is this an application of descriptive or inferential statistics? Explain.

Construct measures of political participation at the nominal, ordinal, and interval-ratio levels. (*Hint:* You can use behaviors such as voting frequency or political party membership.) Discuss the advantages and disadvantages of each.

Variables can be measured according to more than one level of measurement. For the following variables, identify at least two levels of measurement. Is one level of measurement better than another? Explain.

Individual age

Annual income

Religiosity

Student performance

Social class

Number of children

THE ORGANIZATION AND GRAPHIC PRESENTATION OF DATA

Demographers examine the size, composition, and distribution of human populations. Changes in the birth, death, and migration rates of a population affect its composition and social characteristics.¹ To examine a large population, researchers often have to deal with very large amounts of data. For example, imagine the amount of data it takes to describe the immigrant or elderly population in the United States. To make sense out of these data, a researcher must organize and summarize the data in some systematic fashion. In this chapter, we review three such methods used by social scientists: (1) the creation of frequency distributions, (2) the construction of bivariate tables and (3) the use of graphic presentation.

FREQUENCY DISTRIBUTIONS

The most basic way to organize data is to classify the observations into a frequency distribution. A **frequency distribution** is a table that reports the number of observations that fall into each category of the variable we are analyzing. Constructing a frequency distribution is usually the first step in the statistical analysis of data.

Immigration has been described as “remaking America with political, economic, and cultural ramifications.”² Globalization has fueled migration, particularly since the beginning of the 21st century. Workers migrate because of the promise of employment and higher standards of living than what is attainable in their home countries. Data reveal that many migrants seek specifically to move to the United States.³ The U.S. Census Bureau uses the term foreign born to refer to those who are not U.S. citizens at birth. The U.S. Census estimates that 13.5% of the U.S. population, or approximately 44 million people, are foreign born.⁴ Immigrants

Construct and analyze frequency, percentage, and cumulative distributions.

Calculate proportions and percentages.

Compare and contrast frequency and percentage distributions for nominal, ordinal, and interval-ratio variables.

Create a bivariate table

Construct and interpret a pie chart, bar graph, histogram, the statistical map, line graph, and time-series chart.

A table reporting the number of observations falling into each category of the variable.

are not one homogeneous group but are many diverse groups. Table 2.1 shows the frequency distribution of the world region of birth for the foreign-born population.

The frequency distribution is organized in a table, which has a number (2.1) and a descriptive title. The title indicates the kind of data presented: “Frequency Distribution for Categories of Region of Birth for Foreign-Born Population.” The table consists of two columns. The first column identifies the variable (*world region of birth*) and its categories. The second column, with the heading “Frequency (*f*),” tells the number of cases in each category as well as the total number of cases ($N = 43,681,654$). Note also that the source of the table is clearly identified. It tells us that the data are from a 2018 report by Jynnah Radford and Abby Budiman (although the information is based on 2016 American Community Survey data from the U.S. Census). The source of the data can be reported as a source note or in the title of the table. This table is also referred to as a **univariate frequency table**, as it presents frequency information for a single variable.

What can you learn from the information presented in Table 2.1? The table shows that as of 2016, approximately 44 million people were classified as foreign born. Out of this group, most—about 11.7 million people—were from South and East Asia, just under 11.6 million were from Mexico, followed by about 5.8 million from Europe or Canada.

Univariate frequency table: A table that displays the distribution of one variable

Table 2.1 Frequency Distribution for Categories of Region of Birth for Foreign-Born Population, 2016	
Region of Birth	Frequency (<i>f</i>)
South and East Asia	11,731,584
Mexico	11,568,060
Europe/Canada	5,785,135
Caribbean	4,300,022
Central America	3,463,389
South America	2,927,145
Middle East	1,875,264
Sub-Saharan Africa	1,769,778
All other	261,277
Total	43,681,654

Source: “2016, Foreign-Born Population in the United States Statistical Portrait”, Pew Research Center, Washington, D.C (September 14, 2018), <https://www.pewhispanic.org/2018/09/14/2016-statistical-information-on-foreign-born-in-united-states/>.

PROPORTIONS AND PERCENTAGES

Frequency distributions are helpful in presenting information in a compact form. However, when the number of cases is large, the frequencies may be difficult to grasp. To standardize these raw frequencies, we can translate them into relative frequencies—that is, proportions or percentages.

A **proportion** is a relative frequency obtained by dividing the frequency in each category by the total number of cases. To find a proportion (p), divide the frequency (f) in each category by the total number of cases (N):

$$p = \frac{f}{N} \quad (2.1)$$

where

f = frequency

N = total number of cases

We've calculated the proportion for the three largest groups of foreign born. First, the proportion of foreign born originally from South and East Asia is

$$\frac{11,731,584}{43,681,654} = .269$$

The proportion of foreign born who were originally from Mexico is

$$\frac{11,568,060}{43,681,654} = .265$$

The proportion of foreign born who were originally from Europe or Canada is

$$\frac{5,785,135}{43,681,654} = .132$$

The proportion of foreign born who were originally from all other reported areas (combining the category "All other" with those from the Caribbean, Central and South America, Middle East, and sub-Saharan Africa) is

$$\frac{14,596,875}{43,681,654} = .334$$

Proportions should always sum to 1.00 (allowing for some rounding errors). Thus, in our example, the sum of the six proportions is

$$0.27 + 0.27 + 0.13 + 0.33 = 1.0$$

● **Proportion:** A relative frequency obtained by dividing the frequency in each category by the total number of cases.

To determine a frequency from a proportion, we simply multiply the proportion by the total N :

$$f = p(N) \quad (2.2)$$

Thus, the frequency of foreign born from South and East Asia can be calculated as

$$0.27 (43,681,654) = 11,794,047$$

The obtained frequency differs somewhat from the actual frequency of 11,731,584. This difference is due to rounding off of the proportion. If we use the actual proportion instead of the rounded proportion, we obtain the correct frequency:

$$0.268570050026036 (43,681,654) = 11,731,584$$

Percentage: A relative frequency obtained by dividing the frequency in each category by the total number of cases and multiplying by 100.

We can also express frequencies as percentages. A **percentage** is a relative frequency obtained by dividing the frequency in each category by the total number of cases and multiplying by 100. In most statistical reports, frequencies are presented as percentages rather than proportions. Percentages express the size of the frequencies as if there were a total of 100 cases.

To calculate a percentage, multiply the proportion by 100:

$$\text{Percentage (\%)} = \frac{f}{N}(100) \quad (2.3)$$

or

$$\text{Percentage (\%)} = p(100) \quad (2.4)$$

Thus, the percentage of respondents who were originally from Mexico is

$$0.27 (100) = 27\%$$



LEARNING CHECK 2.1

Calculate the proportion and percentage of males and females in your statistics class. What proportion is female?

Percentage distribution: A table showing the percentage of observations falling into each category of the variable.

PERCENTAGE DISTRIBUTIONS

Percentages are usually displayed as percentage distributions. A **percentage distribution** is a table showing the percentage of observations falling into each

Table 2.2 Frequency Distribution for Categories of Region of Birth for Foreign-Born Population, 2016

Region of Birth	Frequency (<i>f</i>)	Percentage (%)
South and East Asia	11,731,584	27
Mexico	11,568,060	27
Europe/Canada	5,785,135	13
Caribbean	4,300,022	10
Central America	3,463,389	8
South America	2,927,145	7
Middle East	1,875,264	4
Sub-Saharan Africa	1,769,778	4
All other	261,277	1
Total	43,681,654	100* (rounded)

Source: “2016, Foreign-Born Population in the United States Statistical Portrait”, Pew Research Center, Washington, DC (September 14, 2018), <https://www.pewhispanic.org/2018/09/14/2016-statistical-information-on-foreign-born-in-united-states/>.

category of the variable. For example, Table 2.2 presents the frequency distribution of categories of places of origin (Table 2.1) along with the corresponding percentage distribution. Percentage distributions (or proportions) should always show the base (N) on which they were computed. Thus, in Table 2.2, the base on which the percentages were computed is $N = 43,681,654$.

THE CONSTRUCTION OF FREQUENCY DISTRIBUTIONS

In this section, you will learn how to construct frequency distributions. Most often, we can use statistical software to accomplish this, but it is important to go through the process to understand how frequency distributions are actually constructed.

For nominal and ordinal variables, constructing a frequency distribution is quite simple. To do so, count and report the number of cases that fall into each category of the variable along with the total number of cases (N). For the purpose of illustration, let's take a small random sample of 40 cases from a General Social Survey (GSS) sample and record their scores on the following variables: gender, a nominal-level variable; degree, an ordinal measurement of education; and age and number of children, both interval-ratio variables. The use of “male” and “female” in parts of this book is in keeping with the GSS categories for the variable *sex* (respondent's sex).